

04/09/2026

Modèle linéaire

C'est comme du cours
aujourd'hui!

$\mathcal{D}_n = (x_i, y_i)_{1 \leq i \leq n}$, $x_i \in \mathbb{R}^p$, y_i : output réel $\in \mathbb{R}$.

On souhaite faire de la régression: trouver 1 fonction f
telle que $y \simeq f(x)$, f est calibrée à partir du jeu
de données $\mathcal{D}_n$.

Stratégie 1: utiliser géom + optim pour trouver f .

Cs Méthode des moindres carrés:

- on met 1 structure sur f , par exemple un mod.
linéaire: $f_\beta(x) = \langle \beta, x \rangle$

- on définit un critère d'être "le meilleur" f pour
approximer y : $\sum_{i=1}^n (y_i - f_\beta(x_i))^2$ est minimal.

Stratégie 2: On met un modèle statistique.

$y = \langle x, \beta \rangle + \sigma \varepsilon$ avec β et σ des
paramètres du modèle. On cherche alors à estimer β et σ
via les obs. $\mathcal{D}_n$. $\varepsilon \sim \mathcal{N}(0, 1)$, y_i indépendantes

$$\left(\mathbb{R}^n, \mathcal{B}(\mathbb{R}^n), \left\{ p_\theta \cdot \text{dLeb}^{\otimes n} := \bigotimes_{i=1}^n \mathcal{N}(\beta_1 + \beta_2 x_i, \sigma^2) : \theta := (\beta_1, \beta_2, \sigma^2) \in \mathbb{R}^2 \times \mathbb{R}_+^* \right\} \right).$$

$\downarrow$

$y_i \text{ sont } \perp.$

$\langle \begin{pmatrix} \beta_1 \\ \beta_2 \end{pmatrix}, \begin{pmatrix} 1 \\ x_i \end{pmatrix} \rangle + \sigma \varepsilon$

Sci, on a $p=1$ et $y_i \simeq \beta_1 + \beta_2 x_i + \sigma \varepsilon_i$

$$Y_i = \langle X_i, \beta \rangle + \sigma \varepsilon_i.$$

$$X_i \in \mathbb{R}^p, \beta \in \mathbb{R}^p$$

1. Montrer que la matrice $X^T X$ est inversible.

(Remarque : Puisque $u^T X^T X u \geq 0$ pour tout $u \in \mathbb{R}^2$, on en déduit que la matrice est définie positive.

↳ C'est vrai tout le temps, il n'y a pas besoin que $p = 1$
 $X^T X$ est toujours ≥ 0

$\forall u$: $u^T X^T X u = \|Xu\|^2 \geq 0$. Donc $X^T X$ est ≥ 0

$X^T X$ est de taille 2×2 .

$$X^T X = \begin{bmatrix} n & \langle 1, x \rangle \\ \langle 1, x \rangle & \|x\|^2 \end{bmatrix} = \begin{pmatrix} n & \sum x_i \\ \sum x_i & \sum x_i^2 \end{pmatrix}$$

$$\det(X^T X) = n \sum x_i^2 - \left(\sum x_i\right)^2$$

≥ 0 par Cauchy-Schwarz et le cas d'égalité est lorsque les vecteurs sont colinéaires

On a $\det(\quad) = 0 \iff 1$ et x st col.

$\iff$ tous les x_i vont être égaux

Or, $\exists i \neq j : x_i \neq x_j$. Donc la conclusion.

Rem : $X^T X$ est une matrice de Gramm.

Cette matrice est inversible lorsque X est de rang maximal, p
 (en admettant que $n > p$)

2. Déterminer les estimateurs du maximum de vraisemblance $(\hat{\beta}_1, \hat{\beta}_2, \hat{\sigma}^2)$ de $(\beta_1, \beta_2, \sigma^2)$.

On va étudier le contexte général : $Y_i = \langle X_i, \beta \rangle + \sigma \varepsilon_i$

$$\varepsilon_i \sim N(0, 1)$$

$$\theta = (\beta, \sigma)$$

$$\theta \mapsto L_n(\theta) = \prod_{i=1}^n f_{\theta}(x_i | x_i)$$

$$= \prod_{i=1}^n \frac{1}{\sqrt{2\pi} \sigma} e^{-\frac{(x_i - \langle x_i, \beta \rangle)^2}{2\sigma^2}}$$

$$\ell_n(\theta) = \log L_n(\theta)$$

$$= -\frac{n}{2} \log(2\pi) - \frac{n}{2} \log(\sigma^2) - \frac{\sum (x_i - \langle x_i, \beta \rangle)^2}{2\sigma^2}$$

$$Y = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} \quad X\beta = \begin{pmatrix} \langle x_1, \beta \rangle \\ \vdots \\ \langle x_n, \beta \rangle \end{pmatrix} \quad \text{avec } X = \begin{pmatrix} x_1^{(p)} & \dots & x_1^{(p)} \\ \vdots & & \vdots \\ x_n^{(p)} & \dots & x_n^{(p)} \end{pmatrix}$$

X : c'est la matrice de design du modèle linéaire.

$$\ell_n(\theta) = -\frac{n}{2} \log(2\pi) - \frac{n}{2} \log(\sigma^2) - \frac{\|Y - X\beta\|^2}{2\sigma^2}$$

Supposons β connu: on observe $\sigma^2 \mapsto \ell_n(\theta)$

$$\partial_{\sigma^2} \ell_n(\theta) = -\frac{n}{2\sigma^2} + \frac{\|Y - X\beta\|^2}{2\sigma^4}$$

$$\partial_{\sigma^2} \ell_n = 0 \iff \hat{\sigma^2} = \frac{\|Y - X\beta\|^2}{n}$$

β sera celui qui maximise

Etudions l'équation en β :

$$\ell_n(\theta) = \text{cte}(\sigma) - \frac{\|Y - X\beta\|^2}{2\sigma^2}$$

On doit minimiser $\|Y - X\beta\|^2 := J(\beta)$

$\beta \mapsto J(\beta)$ est $C^\infty(\mathbb{R}^p, \mathbb{R})$

On calcule $\nabla J(\beta)$ et on cherche le ou les points critiques.

On observe en+ que comme $\beta \mapsto X\beta$ parcourt $\text{Vect}(X)$ alors le min est atteint pour $X\beta = \text{Proj}_{\text{Vect}(X)}(Y)$.

$\nabla J(\beta)$ est l'unique vecteur d_β tq

$$J(\beta+h) = J(\beta) + \langle h, d_\beta \rangle + o(\|h\|)$$

$$\begin{aligned} J(\beta+h) &= \|X(\beta+h) - Y\|^2 \\ &= \|X\beta - Y\|^2 + 2 \langle Xh, X\beta - Y \rangle \\ &\quad + \underbrace{\|Xh\|^2}_{o(\|h\|)} \\ &= J(\beta) + 2 \langle h, \underbrace{X^T(X\beta - Y)}_{\nabla J(\beta)} \rangle + o(\|h\|) \end{aligned}$$

On en déduit que

$$\nabla J(\beta) = 2 X^T(X\beta - Y)$$

Pour la 1^{ère} partie: - soit on dérive le gradient
- soit on fait Taylor ordre 2.

$$D^2 J(\beta) = 2 X^T X$$

On observe que si $X^T X$ est sym $\text{def} \geq 0$, alors J est convexe strictement.

$$\text{On résout } \nabla J(\beta) = 0 \iff X^T X \beta - X^T Y = 0$$

$$\iff X^T X \beta = X^T Y$$

Formule à connaître par cœur.

$$\iff \hat{\beta} = (X^T X)^{-1} X^T Y$$

Dans le cadre de l'exercice:

$$(X^T X)^{-1} = \frac{1}{m \sum x_i^2 - (\sum x_i)^2} \begin{pmatrix} \sum x_i^2 & -\sum x_i \\ -\sum x_i & m \end{pmatrix}$$

$$X^T Y = \begin{pmatrix} 1 & \dots & 1 \\ x_1 & \dots & x_m \end{pmatrix} \begin{pmatrix} y_1 \\ \vdots \\ y_m \end{pmatrix} = \begin{pmatrix} \sum y_i \\ \sum x_i y_i \end{pmatrix}$$

$$(X^T X)^{-1} X^T Y = \frac{\begin{pmatrix} \sum x_i^2 \sum y_i - \sum x_i \sum x_i y_i \\ - \sum x_i \sum y_i + n \sum x_i y_i \end{pmatrix}}{n \sum x_i^2 - (\sum x_i)^2}$$

Dans la logique de la régression, on peut toujours se ramener à un nuage de points centrés: $\sum x_i = 0$

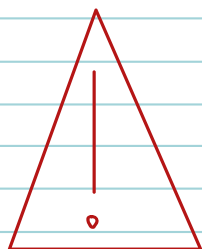
Dans ce cas l'a.c: $\hat{\beta}_1 = \frac{1}{n} \sum y_i$: moyenne des y

$$\hat{\beta}_2 = \frac{\sum x_i y_i}{\sum x_i^2} : \frac{\text{cov des } x \text{ avec } y}{\text{Var}(x)}$$

3. Déterminer la loi de $(\hat{\beta}_1, \hat{\beta}_2)$ sous $p_\theta \cdot d\text{Leb}^{\otimes n}$.

Pour déterminer la loi de $\hat{\beta}$, on va toujours utiliser

$$\hat{\beta} = (X^T X)^{-1} X^T Y$$



Il y a un (β^*, σ^{*2}) qui sert à générer les données.
Je ne les connais pas, et je les cherche.

Il y a aussi $(\hat{\beta}, \hat{\sigma}^2)$ et je cherche leurs lois.

$$Y = X\beta^* + \sigma^* \underbrace{N(0, I_m)}_{\xi}$$

$$\begin{aligned} \hat{\beta} &= (X^T X)^{-1} X^T (X\beta^* + \sigma^* \xi) \\ &= \beta^* + \sigma^* (X^T X)^{-1} X^T \xi \end{aligned}$$

$$\text{Cov}(\sigma^* (X^T X)^{-1} X^T \xi) = \sigma^{*2} \mathbb{E} \left((X^T X)^{-1} X^T \xi \cdot \xi^T X (X^T X)^{-1} \right)$$

$$\begin{aligned} \text{Cov}(A\xi) &= \mathbb{E}(A\xi \xi^T A^T) = \sigma^{*2} (X^T X)^{-1} X^T \underbrace{\mathbb{E}[\xi \xi^T]}_{I_m} X (X^T X)^{-1} \\ &= \sigma^{*2} (X^T X)^{-1} \end{aligned}$$

$$\hat{\beta} \sim N(\beta^*, \sigma^{*2} (X^T X)^{-1})$$

$$E[\hat{\beta}] = \beta^* \quad \text{est sans biais.}$$

$$E \|\hat{\beta} - \beta^*\|^2 = \sigma^{*2} \text{Tr}((X^T X)^{-1})$$

4. Montrer que, sous $p_\theta \cdot d\text{Leb}^{\otimes n}$, $\hat{\beta}_1$ et $\hat{\beta}_2$ sont indépendants si et seulement si $n^{-1} \sum_{i=1}^n x_i =$

On sait que $\hat{\beta}$ est un vecteur Gaussien.

$$\hat{\beta}_1 \perp \hat{\beta}_2 \Leftrightarrow \text{Cov}(\hat{\beta}_1, \hat{\beta}_2) = 0$$

$$\Leftrightarrow \sum x_i = 0 \quad (\text{voir expression de } (X^T X)^{-1})$$

5. Déterminer la loi de $\hat{\sigma}^2$ sous $p_\theta \cdot d\text{Leb}^{\otimes n}$. On sait d'après le 2°) que:

$$\hat{\sigma}^2 = \frac{\|Y - X\hat{\beta}\|^2}{n} = \frac{\|(X\beta^* + \sigma^* \xi) - X(X^T X)^{-1} X^T (X\beta^* + \sigma^* \xi)\|^2}{n}$$

On a remplacé $\hat{\beta}$ par sa valeur en fonction de β^* , σ^* et ξ idem avec Y .

$$\hat{\sigma}^2 = \frac{1}{n} \left\| \cancel{X\beta^*} + \sigma^* \xi - \cancel{X(X^T X)^{-1} X^T X\beta^*} - \cancel{X(X^T X)^{-1} X^T \sigma^* \xi} \right\|^2$$

$$= \frac{\sigma^{*2}}{n} \left\| \xi - X(X^T X)^{-1} X^T \xi \right\|^2$$

$$= \frac{\sigma^{*2}}{n} \left\| (I - P) \xi \right\|^2 \quad \text{où } P = X(X^T X)^{-1} X^T$$

$$\left\{ \begin{array}{l} P \text{ est une projection: } P^2 = P \\ P \text{ est } \perp : P^T = P \end{array} \right\} \text{ sur Vect}(X)$$

$I - P$ est aussi 1 matrice de proj $\perp$ sur $\text{Vect}(X)$.

$$\|(I - P)\xi\|^2 \sim \chi^2(\text{Rank}(I - P))$$

$$\text{Rank}(P) = p. \quad \text{Rank}(I - P) = n - p.$$

$$\mathcal{L}(\hat{\sigma}^2) = \frac{\sigma^{*2}}{n} \chi^2(n-p)$$

$$\mathbb{E}(\hat{\sigma}^2) = \frac{n-p}{n} \sigma^{*2}$$

L'estimateur sans biais associé est $\frac{1}{n-p} \|Y - X\hat{\beta}\|^2$

6. Montrer que sous $p_\theta \cdot d\text{Leb}^{\otimes n}$, $\hat{\sigma}^2$ et $(\hat{\beta}_1, \hat{\beta}_2)$ sont indépendants.

$$\text{On le sait grâce à } (\mathbb{I}-P)P = P(\mathbb{I}-P) = 0$$

Exercice 2 (Estimation d'un coefficient de mélange (*)). Soit un n -échantillon $(X_1, \dots, X_n)$ du modèle $(\mathbb{R}, \mathcal{B}(\mathbb{R}), \{f_\theta \cdot \text{Leb}, \theta \in]0, 1[\})$ où

$$f_\theta(x) = \frac{\theta}{a} \mathbb{1}_{\{[0,a]\}}(x) + \frac{(1-\theta)}{b} \mathbb{1}_{\{[0,b]\}}(x).$$

On suppose que a et b sont connus et que $0 < a < b$. La fonction de répartition associée à la densité f_θ vaut

$$x \mapsto \begin{cases} 0 & x \leq 0 \\ \left(\frac{\theta}{a} + \frac{(1-\theta)}{b}\right)x & 0 \leq x \leq a \\ \theta + \frac{(1-\theta)}{b}x & a \leq x \leq b \\ 1 & x \geq b \end{cases}$$

1. Exprimer la fonction de vraisemblance à l'aide de $N_a := \sum_{i=1}^n \mathbb{1}_{X_i \in [0,a]}$ le nombre d'observations à valeur dans $[0, a]$.

Les obs st iid

$$L_n(\theta) = \prod_{i=1}^n f_\theta(x_i)$$

$$= \prod_{i=1}^n \left(\frac{\theta}{a} \mathbb{1}_{0 \leq x_i \leq a} + \frac{1-\theta}{b} \mathbb{1}_{a \leq x_i \leq b} \right)$$

$$= \prod_{i=1}^n \left(\left(\frac{\theta}{a} + \frac{1-\theta}{b} \right) \mathbb{1}_{x_i \in [0,a]} + \frac{1-\theta}{b} \mathbb{1}_{x_i \notin [0,a]} \right)$$

$$= \left(\frac{\theta}{a} + \frac{1-\theta}{b} \right)^{N_a} \cdot \left(\frac{1-\theta}{b} \right)^{n-N_a}$$

$$\ell_n(\theta) = \log L_n(\theta) = N_a \log \left(\frac{\theta b + a(1-\theta)}{ab} \right) + (n-N_a) \log \frac{1-\theta}{b}$$

$\hat{\theta} = \arg \max_{\theta} P_n(\theta)$. On obtient alors:

$$N_a \frac{(b-a)}{\theta b + (1-\theta)a} - \frac{(n-N_a)}{1-\theta} = 0$$

$$N_a (b-a) (1-\theta) = (n-N_a) (\theta b + (1-\theta)a)$$

$$N_a (b-a) - (n-N_a)a = \theta ((n-N_a)b - (n-N_a)a + N_a(b-a))$$

$$N_a b - n \cdot a = \theta n (b-a)$$

$$\hat{\theta} = \frac{N_a b - n a}{n (b-a)} \vee 0$$

Calculons le max avec 0.

$$\mathbb{E}[\hat{\theta}] = \frac{b \mathbb{E}N_a - n a}{n (b-a)}$$

On observe que:

$$\mathbb{E} \left(\mathbb{1}_{X_i \in [0; a]} \right) = \theta^* + (1-\theta^*) \cdot \frac{a}{b}$$

$$\begin{aligned} \text{D'où } \mathbb{E}[\hat{\theta}] &= \frac{n b \left(\theta^* + (1-\theta^*) \frac{a}{b} \right) - n a}{n (b-a)} \\ &= \frac{n b \theta^* + \cancel{n a (1-\theta^*)} - \cancel{n a}}{n (b-a)} \\ &= \theta^* \end{aligned}$$

3. En utilisant la méthode des moments (d'ordre 1), proposer un estimateur $\hat{\theta}_n^{(2)}$ de θ .

NOT method of moments

$$\mathbb{E}[X] = \theta \cdot \frac{a}{2} + (1-\theta) \frac{b}{2} = \frac{b}{2} + \theta \left(\frac{a-b}{2} \right)$$

$$\frac{1}{n} \sum x_i \simeq \mathbb{E} X \quad \text{LFGN.}$$

$$\text{On résout} \quad \frac{b}{2} + \theta \frac{a-b}{2} = \frac{\sum x_i}{n}$$

$$\hat{\theta}_n^{(2)} = \left(\frac{2 \sum x_i}{n} - b \right) / (a-b)$$

$$= \left(b - \frac{2 \sum x_i}{n} \right) / b - a$$